

(19)日本国特許庁(JP)

(12)特許公報(B2)

(11)特許番号
特許第7599554号
(P7599554)

(45)発行日 令和6年12月13日(2024.12.13)

(24)登録日 令和6年12月5日(2024.12.5)

(51)国際特許分類

F I

G 0 6 T 13/40 (2011.01)

G 0 6 T 13/40

G 1 0 L 25/63 (2013.01)

G 1 0 L 25/63

請求項の数 16 (全13頁)

(21)出願番号	特願2023-514901(P2023-514901)	(73)特許権者	310021766
(86)(22)出願日	令和3年9月2日(2021.9.2)		株式会社ソニー・インタラクティブエン
(65)公表番号	特表2023-540535(P2023-540535		タテインメント
	A)		東京都港区港南1丁目7番1号
(43)公表日	令和5年9月25日(2023.9.25)	(74)代理人	100105924
(86)国際出願番号	PCT/US2021/048870		弁理士 森下 賢樹
(87)国際公開番号	WO2022/051497	(72)発明者	カウシイク、ラクシュミシュ
(87)国際公開日	令和4年3月10日(2022.3.10)		アメリカ合衆国、カリフォルニア州 9
審査請求日	令和5年4月27日(2023.4.27)		4 4 0 4、サン マテオ、ブリッジボイ
(31)優先権主張番号	63/074,310		ント パークウェイ 2 2 0 7
(32)優先日	令和2年9月3日(2020.9.3)	(72)発明者	クマール、サケット
(33)優先権主張国・地域又は機関			アメリカ合衆国、カリフォルニア州 9
	米国(US)		4 4 0 4、サン マテオ、ブリッジボイ
(31)優先権主張番号	17/396,664		ント パークウェイ 2 2 0 7
(32)優先日	令和3年8月7日(2021.8.7)	審査官	中田 剛史
最終頁に続く		最終頁に続く	

(54)【発明の名称】 テキストと音声を用いたフェイシャルアクションユニットの自動生成によるフェイシャルアニメーション制御

(57)【特許請求の範囲】

【請求項1】

装置であって、
少なくとも1つのプロセッサであって、
コンピュータ化されたアバターの顔の画像を識別すること、
前記アバターに関連する第1のモダリティデータであって、音声を含む前記第1のモ
ダリティデータを識別すること、
前記アバターに関連する第2のモダリティデータであって、テキストを含む前記第2
のモダリティデータを識別すること、
前記アバターの前記顔の前記画像をアニメーション化するのに役立つ、前記第1のモ
ダリティデータに少なくとも部分的に基づく第1の情報及び前記第2のモダリティデータ
に少なくとも部分的に基づく第2の情報を受け取ること、及び
前記第1の情報及び前記第2の情報の両方に従って前記アバターの前記顔をアニメ
ーション化すること、
のための命令で構成された、前記少なくとも1つのプロセッサ、
を含み、
前記命令は少なくとも1つの機械学習（ML）モデルにアクセスして前記アバターの前
記顔をアニメーション化することを実行可能であり、前記MLモデルは前記アバターの前
記顔の画像から導出されたアンカーポイントを受け取ってフェイシャルアクションユニッ
ト（FAU）を生成し、前記音声及び前記テキストを前記FAUに関連付けることを実行

10

20

可能であり、

前記命令は、前記音声から抽出された音素境界を少なくとも部分的に使用して前記音声及び前記テキストを時間的に互いに整合させることを実行可能であり、前記命令は、前記音声から抽出された前記音素境界を使用して感情及び感情の程度を前記音声及び前記テキストと時間において整合させることをさらに実行可能である、前記装置。

【請求項 2】

前記命令は、前記第 1 及び第 2 のモダリティデータから前記感情を導出することを実行可能である、請求項 1 に記載の装置。

【請求項 3】

前記情報は、前記感情から生成された時間で整合された単語レベルの感情確率に少なくとも部分的に基づく、請求項 2 に記載の装置。

10

【請求項 4】

方法であって、

第 1 のテキスト及び第 1 の音声の両方に従って単語を話すアニメーション化される第 1 の顔の画像を生成すること、

前記第 1 のテキスト及び前記第 1 の音声を時間において整合させること、及び

前記第 1 のテキスト及び前記第 1 の音声に従って、第 1 の単語を話す前記第 1 の顔の前記画像をアニメーション化すること、
を含み、

前記第 1 の顔の前記画像は、前記第 1 の音声における単語のスライディングウィンドウを処理することによって少なくとも部分的にアニメーション化され、前記第 1 の音声における単語 1 ~ N は第 1 の感情に関連付けられ、前記第 1 の顔の前記画像は前記第 1 の感情に従ってアニメーション化され、前記第 1 の音声における単語 2 ~ N + 1 は第 2 の感情に関連付けられ、前記第 1 の顔の前記画像は前記第 2 の感情に従ってアニメーション化され、N は 2 よりも大きい整数である、前記方法。

20

【請求項 5】

既知の単語を話すアニメーション化された顔のトレーニングセットを使用して、機械学習 (ML) モデルをトレーニングすること、

前記第 1 のテキスト及び前記第 1 の音声を前記 ML モデルに入力すること、

前記 ML モデルの出力に従って、前記第 1 の顔の前記画像をアニメーション化すること、

前記第 1 のテキストから感情と情緒を検出すること、

前記第 1 のテキストを前記第 1 のテキストを表す音声と整合させて、整合されたテキスト / 音声をレンダリングすること、及び

前記感情、前記情緒、及び前記整合されたテキスト / 音声を前記 ML モデルに入力すること、

を含む、請求項 4 に記載の方法。

30

【請求項 6】

ターゲット感情を前記 ML モデルに入力すること

を含む、請求項 5 に記載の方法。

【請求項 7】

前記 ML モデルからフェイシャルアクションを表す第 1 の確率を受け取ること
を含む、請求項 5 に記載の方法。

40

【請求項 8】

前記 ML モデルから感情を表す第 2 の確率を受け取ること
を含む、請求項 7 に記載の方法。

【請求項 9】

前記第 1 及び第 2 の確率の一方または両方を使用して、フェイシャルアクションユニット (FAU) を確立すること

を含む、請求項 8 に記載の方法。

【請求項 10】

50

前記 F A U に従って前記第 1 の顔の前記画像をアニメーション化することを含む、請求項 9 に記載の方法

【請求項 1 1】

アセンブリであって、
アニメーション化されたコンピュータアバターを提示するように構成された少なくとも 1 つのディスプレイ、
機械学習 (M L) モデルを実行するための命令で構成された少なくとも 1 つのプロセッサであって、前記命令は、
前記アバターが話す音声を示すテキストを受け取ること、
音声を受け取ること、
前記音声から抽出された音素境界を少なくとも部分的に使用して前記テキスト及び前記音声を時間において整合させること、
前記 M L モデルを使用して前記テキスト及び前記音声を処理し、フェイシャルアクションユニット (F A U) を生成すること、及び
前記 F A U に従って前記コンピュータアバターをアニメーション化すること、
に対して実行可能である、前記少なくとも 1 つのプロセッサ、
を含む、前記アセンブリ。

10

【請求項 1 2】

前記命令は、
前記テキストから感情及び情緒を検出すること、
前記テキストを、前記テキストを表す音声に整合させて、整合されたテキスト / 音声をレンダリングすること、及び
前記感情、前記情緒、及び前記整合されたテキスト / 音声を前記 M L モデルに入力すること、に対して実行可能である、請求項 1 1 に記載のアセンブリ。

20

【請求項 1 3】

前記命令は、
ターゲット感情を前記 M L モデルに入力することに対して実行可能である、請求項 1 2 に記載のアセンブリ。

【請求項 1 4】

前記命令は、
前記 M L モデルからフェイシャルアクションを表す第 1 の確率を受け取ること
に対して実行可能である、請求項 1 1 に記載のアセンブリ。

30

【請求項 1 5】

前記命令は、
前記 M L モデルから感情を表す第 2 の確率を受け取ること
に対して実行可能である、請求項 1 4 に記載のアセンブリ。

【請求項 1 6】

前記命令は、
前記第 1 及び第 2 の確率を使用して前記 F A U を確立することに対して実行可能である、請求項 1 5 に記載のアセンブリ。

40

【発明の詳細な説明】

【技術分野】

【0 0 0 1】

本願は、概して、コンピュータシミュレーション及び他のアプリケーションでのテキスト及び音声を使用したフェイシャルアクションユニット (F A U) の自動生成によるフェイシャルアニメーション制御に関する。

【背景技術】

【0 0 0 2】

コンピュータゲームのアバターなどのコンピュータ化された画像の顔は、シミュレーションのプレイ中にアニメーション化され、リアルな効果が得られる。本明細書で理解され

50

るように、アーティストは通常、人間の顔のそれぞれの筋肉点の収縮または弛緩をエミュレートすることによって、フレームごとに顔のそれぞれの部分をアニメーション化するために使用されるフェイシャルアニメーションユニット（FAU）を苦労して作成する必要がある。

【発明の概要】

【0003】

装置は、少なくとも1つのプロセッサであって、コンピュータ化されたアバターの顔の画像を識別すること、という命令で構成された、少なくとも1つのプロセッサを含む。命令は、機械学習（ML）エンジンへのアバターに関連する第1のモダリティデータを入力すること、アバターに関連する第2モダリティデータをMLエンジンに入力すること、及びアバターの顔の画像をアニメーション化するのに役立つMLエンジンから出力を受け取ることに對して実行可能である。命令はまた、出力に従ってアバターの顔をアニメーション化することに對して実行可能である。

10

【0004】

いくつかの実施形態で、出力は、複数のフェイシャルアクションユニット（FAU）を含み、各FAUは、顔の画像のそれぞれの部分に関係する。非限定的な例で、第1のモダリティデータがテキストを含み、第2のモダリティデータは音声を含む。

【0005】

非限定的な例では、命令は、MLエンジンを使用して、第1及び第2のモダリティデータから感情アクション情報を導出することに對して実行可能である。出力は、感情アクション情報から生成された時間で整合された単語レベルの感情確率に少なくとも部分的に基づき得る。

20

【0006】

別の態様では、方法が、既知の単語を話すアニメーション化された顔のトレーニングセットを使用して、機械学習（ML）モデルをトレーニングすることを含む。方法はまた、第1のテキストに従って単語を話すアニメーション化される第1の顔の画像を生成することを含む。方法は、第1のテキストをMLモデルに入力すること、及びMLモデルの出力に従って、第1のテキストによって示される第1の単語を話す第1の顔の画像をアニメーション化することをさらに含む。

【0007】

例示的な実施形態では、方法は第1のテキストから感情と情緒を検出すること、第1のテキストを第1のテキストを表す音声と整合させて、整合されたテキスト／音声をレンダリングすること、及び感情、情緒、及び整合されたテキスト／音声をMLモデルに入力することを含み得る。方法はまた、ターゲット感情をMLモデルに入力することを含み得る。方法の例示的な実装形態では、フェイシャルアクションを表すMLモダリティから第1の確率を受け取ること、感情を表すMLモデルから第2の確率を受け取ること、及び第1及び第2の確率を使用して、フェイシャルアクションユニット（FAU）を確立することを含み得る。次いで方法は、FAUに従って第1の顔の画像をアニメーション化することを含むことができる。

30

【0008】

別の態様では、アセンブリが、アニメーション化されたコンピュータアバターを提示するように構成された少なくとも1つのディスプレイ、及び機械学習（ML）モデルを実行するための命令で構成された少なくとも1つのプロセッサを含む。命令は、アバターが話す音声を示すテキストを受け取ること、MLモデルを使用してテキストを処理し、フェイシャルアクションユニット（FAU）を生成すること、及びFAUに従ってコンピュータアバターをアニメーション化することに對して実行可能である。

40

【0009】

本願の詳細は、その構造と動作との両方について、添付の図面を参照すると最もよく理解でき、図面において、類似の参照符号は、類似の部分を指す。

【図面の簡単な説明】

50

【 0 0 1 0 】

【図 1】本原理による例を含む例示的なシステムのブロック図である。

【図 2】例示的なコンピュータ化されたアバターの顔のスクリーンショットである。

【図 3】例示的なフローチャート形式で例示的な全体的なロジックを示す。

【図 4】機械学習 (ML) エンジンのトレーニングプロセスを示す。

【図 5】ML エンジンの推論プロセスを示す。

【発明を実施するための形態】

【 0 0 1 1 】

本開示は、概して、限定されることなく、コンピュータゲームネットワークなどの家電 (CE) デバイスネットワークの態様を含むコンピュータエコシステムに関する。本明細書のシステムは、クライアントコンポーネントとサーバコンポーネントとの間でデータが交換され得るように、ネットワークを通じて接続され得るサーバコンポーネント及びクライアントコンポーネントを含み得る。クライアントコンポーネントは、Sony PlayStation (登録商標) などのゲームコンソールまたはMicrosoft (登録商標) もしくはNintendo (登録商標) もしくは他の製造者によって作成されたゲームコンソール、仮想現実 (VR) ヘッドセット、拡張現実 (AR) ヘッドセット、ポータブルテレビ (例えば、スマートテレビ、インターネット対応テレビ)、ラップトップ及びタブレットコンピュータなどのポータブルコンピュータ、ならびにスマートフォン及び以下で議論される追加の実施例を含む他のモバイルデバイスを含む、1つ以上のコンピューティングデバイスを含み得る。これらのクライアントデバイスは、様々な動作環境で動作し得る。例えば、クライアントコンピュータのいくつかは、実施例として、Linux (登録商標) オペレーティングシステム、Microsoft (登録商標) のオペレーティングシステム、またはUnix (登録商標) オペレーティングシステム、またはApple, Inc. (登録商標) もしくはGoogle (登録商標) によって制作されたオペレーティングシステムを採用し得る。これらの動作環境は、Microsoft (登録商標) もしくはGoogle (登録商標) もしくはMozilla (登録商標) によって作成されたブラウザ、または以下で議論されるインターネットサーバによってホストされるウェブサイトにアクセスできる他のブラウザプログラムなど、1つ以上の閲覧プログラムを実行するために使用され得る。また、本原理による動作環境を使用して、1つ以上のコンピュータゲームプログラムを実行し得る。

【 0 0 1 2 】

サーバ及び/またはゲートウェイは、インターネットなどのネットワークを通じてデータを受信及び送信するようにサーバを構成する命令を実行する1つ以上のプロセッサを含み得る。あるいは、クライアント及びサーバは、ローカルイントラネットまたは仮想プライベートネットワークを通じて接続することができる。サーバまたはコントローラは、Sony PlayStation (登録商標) などのゲームコンソール、パーソナルコンピュータなどによってインスタンス化され得る。

【 0 0 1 3 】

クライアントとサーバとの間でネットワークを通じて情報を交換し得る。この目的及びセキュリティのために、サーバ及び/またはクライアントは、ファイアウォール、ロードバランサ、テンポラリストレージ、及びプロキシ、ならびに信頼性及びセキュリティのための他のネットワークインフラストラクチャを含むことができる。1つ以上のサーバは、ネットワークメンバーにオンラインソーシャルウェブサイトなどの安全なコミュニティを提供する方法を実装する装置を形成し得る。

【 0 0 1 4 】

プロセッサは、アドレスライン、データライン及び制御ラインなどの様々なライン、並びにレジスタ及びシフトレジスタによって論理を実行することができる、シングルチッププロセッサまたはマルチチッププロセッサであってよい。

【 0 0 1 5 】

一実施形態に含まれるコンポーネントは、他の実施形態では、任意の適切な組み合わせ

で使用することができる。例えば、本明細書に記載される、及び／または図で示される様々なコンポーネントのいずれもは、組み合わせられ、交換され、または他の実施形態から除外されてもよい。

【 0 0 1 6 】

「 A、 B 及び C のうちの少なくとも 1 つを有するシステム」（同様に「 A、 B または C のうちの少なくとも 1 つを有するシステム」及び「 A、 B、 C のうちの少なくとも 1 つを有するシステム」）は、 A 単独、 B 単独、 C 単独、 A 及び B を一緒に、 A 及び C を一緒に、 B 及び C を一緒に、ならびに／または A、 B 及び C を一緒に有するシステムなどを含む。

【 0 0 1 7 】

ここで、具体的に図 1 を参照すると、本原理による、上述され、以下でさらに説明される例示的なデバイスのうちの 1 つ以上を含み得る例示的なシステム 1 0 が示されている。システム 1 0 に含まれる例示的なデバイスのうちの第 1 のデバイスは、限定されることなく、テレビチューナ（同等に、テレビを制御するセットトップボックス）を備えたインターネット対応テレビなどのオーディオビデオデバイス（ A V D ） 1 2 などの家電（ C E ）デバイスである。代替として、 A V D 1 2 は、また、コンピュータ制御型インターネット対応（「スマート」）電話、タブレットコンピュータ、ノートブックコンピュータ、 H M D、ウェアラブルコンピュータ制御デバイス、コンピュータ制御型インターネット対応ミュージックプレイヤー、コンピュータ制御型インターネット対応ヘッドフォン、インプラント可能な皮膚用デバイスなどのコンピュータ制御型インターネット対応インプラント可能デバイス、などであってもよい。それに関わらず、 A V D 1 2 は、本原理を実施する（例えば、本原理を実施するように他の C E デバイスと通信し、本明細書に記載される論理を実行し、本明細書に記載されるいずれかの他の機能及び／または動作を行う）ように構成されることを理解されたい。

【 0 0 1 8 】

したがって、このような原理を実施するために、 A V D 1 2 は、図 1 に示されているコンポーネントの一部または全てによって確立することができる。例えば、 A V D 1 2 は、 1 つ以上のディスプレイ 1 4 を備えることができ、このディスプレイは、高解像度もしくは超高解像度「 4 K 」またはそれ以上の解像度のフラットスクリーンによって実装されてもよく、ディスプレイのタッチを介したユーザ入力信号を受信するためにタッチ対応であってもよい。 A V D 1 2 は、本原理に従ってオーディオを出力するための 1 つ以上のスピーカ 1 6、及び可聴コマンドを A V D 1 2 に入力して A V D 1 2 を制御するためのオーディオ受信機／マイクロホンなどの、少なくとも 1 つの追加入力デバイス 1 8 を含み得る。例示的な A V D 1 2 は、また、 1 つ以上のプロセッサ 2 4 の制御の下、インターネット、 W A N、 L A N などの少なくとも 1 つのネットワーク 2 2 を通じて通信するための 1 つ以上のネットワークインタフェース 2 0 を含み得る。また、グラフィックプロセッサ 2 4 A が含まれていてもよい。したがって、インタフェース 2 0 は、限定されることなく、 W i - F i （登録商標）送受信機であり得て、この W i - F i （登録商標）送受信機は、限定されることなく、メッシュネットワーク送受信機などの無線コンピュータネットワークインタフェースの実施例である。プロセッサ 2 4 は、その上に画像を提示するようにディスプレイ 1 4 を制御すること及びそこから入力を受信することなど、本明細書に記載される A V D 1 2 の他の要素を含む A V D 1 2 が本原理を実施するように、制御することを理解されたい。さらに、ネットワークインタフェース 2 0 は、有線もしくは無線のモデムもしくはルータ、または、例えば、無線テレフォニ送受信機もしくは上述した W i - F i （登録商標）送受信機などの他の適切なインタフェースであってもよいことに留意されたい。

【 0 0 1 9 】

上記のものに加えて、 A V D 1 2 はまた、例えば、別の C E デバイスに物理的に接続する高解像度マルチメディアインタフェース（ H D M I （登録商標））ポートもしくは U S B ポート、及び／またはヘッドフォンを通して A V D 1 2 からユーザにオーディオを提供するために A V D 1 2 にヘッドフォンを接続するヘッドフォンポートなどの 1 つ以上の入力ポート 2 6 を含んでもよい。例えば、入力ポート 2 6 は、オーディオビデオコンテンツ

10

20

30

40

50

のケーブルまたは衛星ソース 26a に有線または無線で接続されてもよい。したがって、ソース 26a は、別個のもしくは統合されたセットトップボックス、または衛星受信機であってよい。あるいは、ソース 26a は、コンテンツを含むゲームコンソールまたはディスクプレイアであってよい。ソース 26a は、ゲームコンソールとして実装されるとき、CE デバイス 44 に関連して以下で説明されるコンポーネントの一部または全てを含んでよい。

【0020】

AVD 12 は、さらに、一時的信号ではない、ディスクベースストレージまたはソリッドステートストレージなどの 1 つ以上のコンピュータメモリ 28 を含んでもよく、これらのストレージは、場合によっては、スタンドアロンデバイスとして AVD のシャーシ内で、または AV プログラムを再生するために AVD のシャーシの内部もしくは外部のいずれかでパーソナルビデオ録画デバイス (PVR) もしくはビデオディスクプレイアとして、または取り外し可能メモリ媒体として具現化されてもよい。また、ある実施形態では、AVD 12 は、限定されることなく、携帯電話受信機、GPS 受信機、及び/または高度計 30 などの位置または場所の受信機を含むことができ、位置または場所の受信機は、衛星もしくは携帯電話基地局から地理的位置情報を受信し、その情報をプロセッサ 24 に供給し、及び/または AVD 12 がプロセッサ 24 と併せて配置されている高度を決定するように構成される。コンポーネント 30 はまた、通常、加速度計、ジャイロスコープ、及び磁力計の組み合わせを含み、AVD 12 の位置及び方向を 3 次元で決定する慣性測定ユニット (IMU) によって実装されてもよい。

【0021】

AVD 12 の説明を続けると、いくつかの実施形態では、AVD 12 は、1 つ以上のカメラ 32 を含んでもよく、1 つ以上のカメラは、サーマルイメージングカメラ、ウェブカメラなどのデジタルカメラ、及び/または AVD 12 に統合され、本原理に従って写真/画像及び/またはビデオを収集するようプロセッサ 24 によって制御可能なカメラであってよい。また、AVD 12 に含まれるのは、Bluetooth (登録商標) 及び/または近距離無線通信 (NFC) 技術を各々使用して、他のデバイスと通信するための Bluetooth (登録商標) 送受信機 34 及び他の NFC 要素 36 であってよい。例示的な NFC 素子は、無線周波数識別 (RFID) 素子であってよい。

【0022】

さらにまた、AVD 12 は、プロセッサ 24 に入力を供給する 1 つ以上の補助センサ 37 (例えば、加速度計、ジャイロスコープ、サイクロメータなどの運動センサ、または磁気センサ、赤外線 (IR) センサ、光学センサ、速度センサ及び/またはケイデンスセンサ、ジェスチャセンサ (例えば、ジェスチャコマンドを検知するための) など) を含み得る。AVD 12 は、プロセッサ 24 への入力をもたらし OTA (無線) TV 放送を受信するための無線 TV 放送ポート 38 を含み得る。上記に加えて、AVD 12 はまた、赤外線データアソシエーション (IRDA) デバイスなどの赤外線 (IR) 送信機及び/または IR 受信機及び/または IR 送受信機 42 を含み得ることに留意されたい。電池 (図示せず) は、電池を充電するために及び/または AVD 12 に電力を供給するために運動エネルギーを電力に変えることができる運動エネルギーハーベスタのように、AVD 12 に電力を供給するために提供され得る。

【0023】

さらに図 1 を参照すると、AVD 12 に加えて、システム 10 は、1 つ以上の他の CE デバイスタイプを含み得る。一実施例では、第 1 の CE デバイス 44 は、AVD 12 に直接送信されるコマンドを介して及び/または後述のサーバを通して、コンピュータゲームの音声及びビデオを AVD 12 に送信するために使用することができるコンピュータゲームコンソールであり得る一方で、第 2 の CE デバイス 46 は第 1 の CE デバイス 44 と同様のコンポーネントを含み得る。図示の実施例では、第 2 の CE デバイス 46 は、プレイヤによって操作されるコンピュータゲームのコントローラとして、またはプレイヤ 47 によって装着されるヘッドマウントディスプレイ (HMD) として構成され得る。図示の実

施例では、2つのC Eデバイス44、46のみが示されているが、より少ないまたはより多くのデバイスが使用されてよいことは理解されよう。本明細書のデバイスは、A V D 12について示されているコンポーネントの一部またはすべてを実装し得る。次の図に示されているコンポーネントのいずれかに、A V D 12の場合に示されているコンポーネントの一部またはすべてが組み込まれることがある。

【0024】

ここで、上述の少なくとも1つのサーバ50を参照すると、サーバは、少なくとも1つのサーバプロセッサ52と、データベースストレージまたはソリッドステートストレージなどの少なくとも1つの有形コンピュータ可読記憶媒体54と、サーバプロセッサ52の制御下で、ネットワーク22を通じて図1の他のデバイスとの通信を可能にし、実際に、本原理に従ってサーバとクライアントデバイスとの間の通信を容易にし得る少なくとも1つのネットワークインタフェース56とを含む。ネットワークインタフェース56は、例えば、有線もしくは無線モデムもしくはルータ、Wi-Fi送受信機、または、例えば、無線テレフォニ送受信機などの他の適切なインタフェースであってよいことに留意されたい。

10

【0025】

したがって、いくつかの実施形態では、サーバ50は、インターネットサーバまたはサーバ「ファーム」全体であってよく、「クラウド」機能を含んでもよく、「クラウド」機能を実行してもよく、システム10のデバイスが、例えば、ネットワークゲームアプリケーションの例示的な実施形態においてサーバ50を介して「クラウド」環境にアクセスし得るようにする。あるいは、サーバ50は、図1に示されている他のデバイスと同じ部屋にある、またはその近くにある、1つ以上のゲームコンソール、または他のコンピュータによって実装されてもよい。

20

【0026】

図2は、コンピュータアバターの顔202が提示される本明細書の任意のディスプレイなどのディスプレイ200を示す。顔202は、人間の顔の対応する部分の弛緩または収縮をエミュレートするために、それぞれのフェイシャルアニメーションユニット(F A U)によって要求されるようにアニメーション化され得る、多くのエミュレートされたアンカーポイント204(明確にするために4つだけが示されている)を有する。

【0027】

30

図3は、F A Uを自動的に生成して顔202をアニメーション化するために使用できる全体的なロジックを示す。本明細書で説明するロジックは、本明細書で説明するプロセッサのうちの任意の1つまたは複数によって実行することができる。

【0028】

ブロック300で始まり、機械学習(M L)エンジンが、図4を参照して説明されるように訓練される。ブロック302に進むと、顔のビデオ画像202が、例えばコンピュータシミュレーション、例えばコンピュータゲームにより生成され、顔に関連する対応するマルチモダリティデータがブロック304でM Lエンジンに入力される。図5を参照してさらに十分に説明するように、マルチモダリティデータから、M Lエンジンは、ブロック306でアバターの顔202のエミュレートされた筋肉をアニメーション化するために使用される予測F A Uを生成する。

40

【0029】

図4は、M Lエンジンのトレーニングプロセスの例を示している。データベース400からのデータのトレーニングセットがアクセスされる。データのトレーニングセットには、話しているときに感情を表現する顔の画像が含まれている。換言すれば、データベースは、正しい表現をアニメーション化しながら既知のテキストを話すアバターのグラウンドトゥールスアニメーションを含む。

【0030】

入力されるトレーニング画像ごとに、画像のそれぞれのアンカーポイント204がブロック402で画像から抽出される。言い換えれば、ブロック402は、フェイシャルアク

50

ションユニット（FAU）に使用される入力画像のアンカーポイントを生成する。また、音声信号及び音声信号が関連付けられているテキストは、ブロック404で抽出され、ブロック404での音声及びテキストがブロック402で抽出されたFAUに関連付けられていることが理解される。

【0031】

ブロック406は、以下にさらに開示するように、音声とテキストが互いに時間的に整合していることを示す。図4のグラフ407によって表されるように、音素境界は、対応するテキスト407Bの音声信号407Aの音素境界から抽出される。このようにすると、x軸に沿った各インクリメンタルブロックが音声によって生成された音素を表し、音素がインクリメンタルブロックの隣接する境界407C、407Dの間で定められる。

10

【0032】

強制整合ブロック406は、音素境界を使用して、感情注釈ブロック408で感情を時間において整合し、エキスパートが感情で入力に注釈を付ける感情注釈ブロック410で感情を整合し、音声特徴ブロック412で音声の特徴を整合させる。したがって、対応する感情、情緒、及び音声の特徴は、抽出された音素境界を使用して整合される。以下でさらに説明するように、「情緒」は感情の程度、例えば肯定的または否定的であることを指すことができ、「感情」は話者の感情的な状態、例えば幸福、悲しみ、怒りなどを指すことができる。音声における単語のスライディングウィンドウは、この方法で処理できる、例えば、単語1～3に第1の感情/情緒で注釈を付けることができ、単語2～4に第2の感情/情緒で注釈を付けることができる、等々であることに留意されたい。音声特徴412は、入力の組み合わせられた特性を表すために生成される。ウィグナー・ヴィル分布関数及び時間周波数フィルタリング技法を含む技法を使用して音声特徴を生成することができる。

20

【0033】

ブロック408、410、及び412の生成物は、ML感情アクション生成モデル414に供給され、モデルをトレーニングする。また、モデル414への入力、ブロック402からのFAUと、モデル414をトレーニングするために使用される入力される（注釈付き）感情のグラウンドトゥールズであるターゲット感情416である。このようにすると、モデル414はデータベース400からの複数のデータセットでトレーニングされ、FAUに対応する音声及びテキスト及びターゲット感情と関連させて、フェイシャルアクションアンカー確率418及び時間で整合させた単語レベル感情確率420を出力し、その後、FAUに従ってアンカーポイント204を適切にアニメーション化することによって、顔（例えば、図2の顔202）をアニメーション化することができる。

30

【0034】

図5に移る前に、感情は、カテゴリ及び次元な感情の分類の一方または両方を使用してグラウンドトゥールズ416と比較するために、入力トレーニングセットで検出され得る。カテゴリ分類は、幸福、悲しみ、怒り、恐怖などを含む感情カテゴリを生成する。次元な分類は、感情価と覚醒度の観点から感情の次元を生成する。

【0035】

図5は、モデル414によるその後の推論処理を示し、図5では510とラベル付けされている。アニメ化されたアバターが話すテキスト500は、テキストを音声信号に変換するためにテキスト読み上げブロック502に入力される。さらに、または代わりに、テキストに対応する記録された、または生の人間の音声ブロック502に供給され得る。図4を参照して上で教示されたように整合されたものに対応するブロックの要素を時間的に整合させるために、図4のグラフ407と同じ特徴を有するグラフ507によって示されるように、図4のトレーニングプロセスについて説明されたのと同じ方法で、ブロック504で音声の特徴が抽出され、ブロック506でテキスト及び対応する音声から音素境界が抽出される。

40

【0036】

ブロック500からのアバターが話すテキストは感情及び情緒検出ブロック508に送

50

られ、図4からのトレーニングに従ってテキストから情緒及び感情を抽出する。ブロック508でテキストから抽出された感情及び情緒は、ブロック504からの音声の特徴及びブロック506からの整合されたテキスト/音声信号とともに、モデル510に入力される。所望であれば、ブロック512に示されるように、ターゲット感情がモデルに入力され得る。ターゲット感情は、ユーザが入力するか、機械学習を使用して入力されたテキストから導出される、注釈付きテキストの一部であってもよい。

【0037】

そのトレーニングに基づいて、モデル510は、受け取った入力から、各アンカーポイント204のフェイシャルアクションの確率512と、時間で整合されたテキスト/音声信号から導出された感情の確率516を出力する。これらの確率の一方または両方に基づいて、投影されたアンカー位置518（すなわち、FAU）が得られ、各アニメーション化されたフレームが描写する顔の表情を確立する。ブロック500で、FAUがアバターのアニメーションに適用され、アバターがテキスト入力を話すときに正しい表情を与える。

【0038】

いくつかの例示的な実施形態を参照して本原理を説明したが、これらは限定することを意図しておらず、各種の代替的な構成が本明細書で特許請求される主題を実施するために使用されてよいことは理解されよう。

10

20

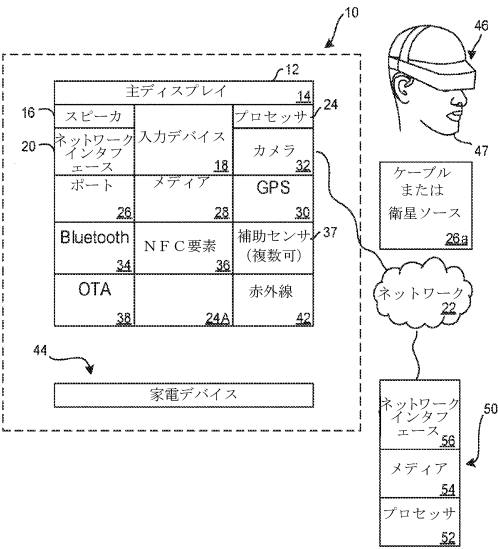
30

40

50

【図面】

【図 1】



【図 2】

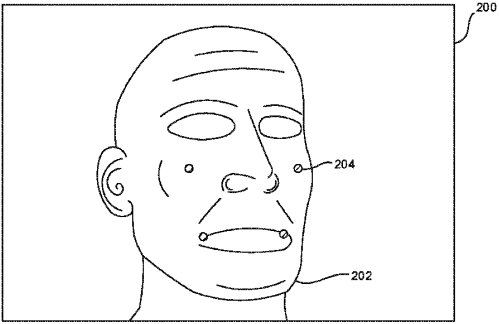
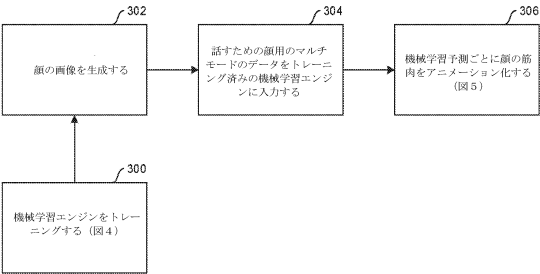
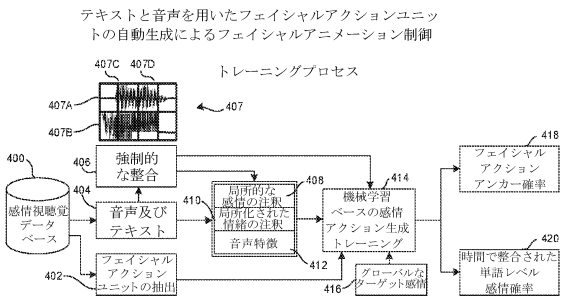


FIG. 2

【図 3】



【図 4】



10

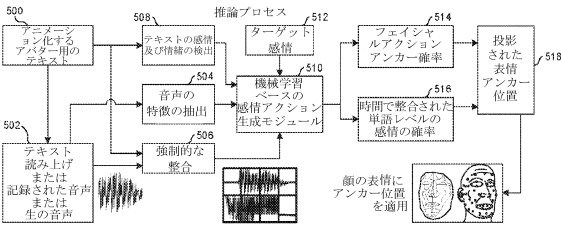
20

30

40

50

【図 5】



10

20

30

40

50

フロントページの続き

- (33)優先権主張国・地域又は機関
米国(US)
- (56)参考文献 国際公開第 2 0 1 9 / 1 6 0 1 0 0 (W O , A 1)
特開 2 0 0 7 - 0 5 8 8 4 6 (J P , A)
特開 2 0 0 5 - 3 5 2 8 9 2 (J P , A)
- (58)調査した分野 (Int.Cl. , D B 名)
G 0 6 T 1 3 / 4 0
G 1 0 L 2 5 / 6 3